

Table of Contents

Benchmarking and Evaluation Initiatives

Analysis and Refinement of Cross-Lingual Entity Linking	1
<i>Taylor Cassidy, Heng Ji, Hongbo Deng, Jing Zheng, and Jiawei Han</i>	
Seven Years of INEX Interactive Retrieval Experiments – Lessons and Challenges	13
<i>Ragnar Nordlie and Nils Pharo</i>	
Bringing the Algorithms to the Data: Cloud-Based Benchmarking for Medical Image Analysis	24
<i>Allan Hanbury, Henning Müller, Georg Langs, Marc André Weber, Bjoern H. Menze, and Tomas Salas Fernandez</i>	
Going beyond CLEF-IP: The ‘Reality’ for Patent Searchers?	30
<i>Julia J. Jürgens, Preben Hansen, and Christa Womser-Hacker</i>	
MusiClef: Multimodal Music Tagging Task	36
<i>Nicola Orio, Cynthia C.S. Liem, Geoffroy Peeters, and Markus Schedl</i>	

Information Access

Generating Pseudo Test Collections for Learning to Rank Scientific Articles	42
<i>Richard Berendsen, Manos Tsagkias, Maarten de Rijke, and Edgar Meij</i>	
Effects of Language and Topic Size in Patent IR: An Empirical Study	54
<i>Florina Piroi, Mihai Lupu, and Allan Hanbury</i>	
Cross-Language High Similarity Search Using a Conceptual Thesaurus	67
<i>Parth Gupta, Alberto Barrón-Cedeño, and Paolo Rosso</i>	
The Appearance of the Giant Component in Descriptor Graphs and Its Application for Descriptor Selection	76
<i>Anita Keszler, Levente Kovács, and Tamás Szirányi</i>	
Hidden Markov Model for Term Weighting in Verbose Queries	82
<i>Xueliang Yan, Guanglai Gao, Xiangdong Su, Hongxi Wei, Xueliang Zhang, and Qianqian Lu</i>	

Evaluation Methodologies and Infrastructure

DIRECTIONS: Design and Specification of an IR Evaluation Infrastructure 88
 Maristella Agosti, Emanuele Di Buccio, Nicola Ferro, Ivano Masiero, Simone Peruzzo, and Gianmaria Silvello

Penalty Functions for Evaluation Measures of Unsegmented Speech Retrieval 100
 Petra Galuščáková, Pavel Pecina, and Jan Hajič

Cumulated Relative Position: A Metric for Ranking Evaluation 112
 Marco Angelini, Nicola Ferro, Kalervo Järvelin, Heikki Keskustalo, Ari Pirkola, Giuseppe Santucci, and Gianmaria Silvello

Better than Their Reputation? On the Reliability of Relevance Assessments with Students 124
 Philipp Schaer

Posters

Comparing IR System Components Using Beanplots 136
 Jens Kürsten and Maximilian Eibl

Language Independent Query Focused Snippet Generation 138
 Pinaki Bhaskar and Sivaji Bandyopadhyay

A Test Collection to Evaluate Plagiarism by Missing or Incorrect References 141
 Solange de L. Pertile and Viviane P. Moreira

Author Index 145